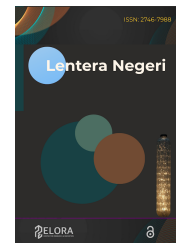




Contents lists available at [Elora Center](#)

Lentera Negeri

Journal homepage: <https://lentera.eloracenter.org/lentera>



Artificial intelligence integration and self-regulated learning in English as a foreign language: Scaffold or substitute

Yuliana Friska^{1,2}, Widodo Widodo¹, Merry Lapasau¹, Syahfitri Purnama¹

¹ Education Department of Postgraduate Faculty, Universitas Indraprasta PGRI, Jakarta Indonesia

² Universitas Pamulang, Indonesia

Article Info

Article history:

Received Jun 12th, 2026

Revised Jul 01th, 2026

Accepted Jul 08th, 2026

Keyword:

Artificial intelligence,
Self-regulated learning,
Systematic review,
PRISMA 2020,
Generative AI

ABSTRACT

Artificial intelligence integration has increasingly been positioned as a mechanism for supporting self-regulated learning (SRL) among English as a Foreign Language (EFL) learners, yet the empirical evidence remains scattered across tool types and contexts. This systematic review synthesized 22 peer-reviewed studies published between 2020 and 2026, comprising 13 experimental or quasi-experimental, 6 mixed-methods, and 3 qualitative designs, examining what AI tools and frameworks characterize the literature, how AI integration affects SRL dimensions, and what factors moderate the AI-SRL relationship. Following PRISMA 2020 guidelines, studies meeting seven inclusion criteria, spanning population, design, document type, language, timeframe, and validated SRL measurement, were retrieved from Scopus and Web of Science, screened with substantial inter-rater agreement ($\kappa = .83$ abstract stage, $.79$ full-text), and appraised using the Mixed Methods Appraisal Tool, whose single framework accommodates all three design types. The AI tool landscape has shifted markedly. Generative AI platforms, particularly ChatGPT and its derivatives, now account for 64% of studies, compared to 18% for earlier automated writing evaluation tools. AI integration most consistently supported goal-setting, self-monitoring, and metacognitive strategy use, while motivational regulation proved less stable, with five of 22 studies reporting reduced autonomous motivation among high-AI-use learners. AI functions as an effective regulatory scaffold within structured pedagogical sequences, though benefits depend on tool type, learner profile, and EFL context. Future research should examine motivational regulation and metacognitive cycles longitudinally, particularly in underrepresented Asian EFL settings.



© 2026 The Authors.

This is an open access article under the CC BY-NC-SA license
(<https://creativecommons.org/licenses/by-nc-sa/4.0>)

Corresponding Author:

Yuliana Friska,

yulianafriska22@gmail.com

Introduction

Research on AI integration in EFL contexts has expanded rapidly since generative AI platforms became widely accessible. ChatGPT represents the most consequential entry point in this shift (Kasneji et al., 2023). Much of this research converges on self-regulated learning (SRL), treating it not as a background variable but as the primary mechanism through which technology access produces measurable learning gains. SRL describes learners who set goals deliberately, monitor their own performance, select appropriate task strategies, and reflect on outcomes to adjust subsequent learning cycles (Pintrich, 2004; Zimmerman, 2000).

In EFL contexts, where authentic second language (L2) exposure outside the classroom remains limited, SRL has consistently predicted achievement, persistence, and long-term proficiency development (Habók & Magyar, 2022; Ismail et al., 2023). The integration of AI into these environments raises a question that existing frameworks have not yet fully addressed, does AI integration strengthen learners' capacity to self-regulate, or does it risk displacing the cognitive effort that regulation fundamentally requires? Hu and Zhang (2025) and Tran et al. (2025) come closest to answering this question empirically, each documenting AI-induced motivational displacement within a single theoretical frame and a single national context. Whether this displacement generalizes across settings, or reflects an artifact of specific instructional conditions, remains unresolved.

Despite a growing empirical base, the literature remains fragmented. Studies differ considerably in how SRL is defined and operationalized. Some adopt Zimmerman's three-phase cyclical model, others rely on Pintrich's componential Motivated Strategies for Learning Questionnaire (MSLQ) framework, and several employ hybrid instruments combining constructs from Self-Determination Theory, Achievement Goal Theory, or Control-Value Theory. The AI tools under investigation are equally varied, spanning automated writing evaluation (AWE) systems, generative AI chatbots, adaptive reading platforms, and speech-focused applications. Without systematic synthesis, the field cannot offer a coherent account of which AI applications most reliably support SRL outcomes, under what conditions, and across which skill domains. A synthesis approach, rather than additional primary data collection, is warranted here precisely because the fragmentation problem is cross-study. No single empirical investigation can resolve disagreements that exist between studies using incompatible SRL operationalizations, and only systematic aggregation following an explicit protocol, as outlined by PRISMA (Page et al., 2021), can establish which patterns replicate across contexts and which are artifacts of a single research design.

Prior syntheses have examined AI in language learning broadly (Zawacki-richter et al., 2019) or SRL in digital environments in general terms (Broadbent & Poon, 2015). Chiu (2024) offers a conceptual analysis of generative AI's educational implications, though without specific attention to SRL outcomes or EFL contexts. Neither prior synthesis, in other words, isolated SRL as a distinct outcome variable while also treating AI tool type as a differentiating factor. Crucially, no existing synthesis has treated AI tool type as a theoretically active variable while also operationalizing SRL as a measurable outcome and applying systematic quality appraisal, all within EFL settings. This is not simply a coverage gap. SDT, CVT, AGT, and Zimmerman's model generate conflicting predictions about when AI functions as scaffold rather than substitute, as detailed in the Theoretical Frameworks section below, and no single theoretical lens can resolve that tension alone. Synthesis becomes necessary because these frameworks diverge, not because the field lacks studies that cover AI, SRL, and EFL simultaneously. Beyond this methodological gap, the EFL context itself remains undertheorized. Restricted authentic input, examination-driven motivational orientations, and deeply socialized patterns of teacher-directive instruction are not incidental features of EFL learning. They fundamentally shape whether AI tools function as regulatory supports or, alternatively, as substitutes that erode the very cognitive effort that self-regulation requires.

The present study addresses these gaps by applying PRISMA 2020 guidelines (Page et al., 2021) alongside the MMAT (Hong et al., 2018), producing a quality-weighted evidence base rather than one organized by publication volume alone. Three research questions guide the analysis. RQ1: What types of AI tools and SRL theoretical frameworks characterize the empirical literature on AI-integrated self-regulated learning in EFL contexts from 2020 to 2026? RQ2: How does AI integration affect EFL learners' self-regulated learning, and which SRL dimensions show the most consistent outcomes? RQ3: What contextual and methodological factors moderate the relationship between AI integration and SRL development in EFL settings?

Self-Regulated Learning: Theoretical Frameworks

SRL describes the process by which learners direct their cognitive, metacognitive, motivational, and behavioral resources toward academic goals (Zimmerman, 2000). The corpus examined in this review draws on at least six distinct theoretical frameworks, a diversity that reflects genuine theoretical richness while also limiting cross-study comparability. Two frameworks dominate. Zimmerman's three-phase cyclical model covers forethought, performance, and self-reflection. Pintrich's componential MSLQ framework partitions SRL into motivational beliefs, cognitive strategies, metacognitive regulation, and resource management (Pintrich, 2004).

Zimmerman's (2000) cyclical model is directly relevant to AI-integrated contexts, as each feedback interaction with an AI system constitutes a micro-cycle that either strengthens or disrupts the learner's self-regulatory loop. Where Zimmerman's model captures the temporal sequence of learning cycles, MSLQ-

based structure is more sensitive to differential effects (Pintrich, 2004). An AI tool might strengthen metacognitive regulation while simultaneously eroding effort regulation if learners outsource persistence to the system. Studies using the MSLQ are accordingly more likely to detect component-level variation in AI effects. The metacognitive regulation component, rooted in foundational work on monitoring and control, emerged as the SRL dimension most consistently activated by AI across the included corpus (Teng, 2021).

These two frameworks account for 59% of the corpus. Beyond Zimmerman and Pintrich, the corpus also draws on Self-Determination Theory (SDT) (Ryan & Deci, 2000), which identifies autonomy, competence, and relatedness as the needs underlying intrinsic motivation and autonomous self-regulation. Izadpanah's (2022) SDL model foregrounds learner control and self-management, while Control-Value Theory (Pekrun, 2024) links academic emotions to motivation and engagement. Achievement Goal Theory (Chazan et al., 2022) further distinguishes mastery-approach from performance-approach orientations. Each framework generates distinct predictions about how AI tools interact with self-regulatory processes. In this review, theoretical plurality is treated as an empirical variable rather than a background condition.

These frameworks do not converge on a single prediction about AI's regulatory function. SDT (Ryan & Deci, 2000) predicts that AI-integrated learning erodes autonomous motivation only when AI use is externally mandated rather than volitionally adopted, since need-satisfaction, not tool exposure, drives self-regulation. CVT locates the mechanism elsewhere. Pekrun (2024) argued that AI's effect on SRL is mediated by the achievement emotions AI feedback elicits, so that immediate feedback strengthens engagement when it produces positive emotion but disrupts regulation when perceived as evaluative and anxiety-inducing. AGT diverges again. Whether learners adopt AI-assisted learning under a mastery-approach or a performance-approach orientation determines its regulatory outcome, and only the former is likely to support the internalization scaffolding requires (Chazan et al., 2022). Three frameworks, three distinct mechanisms, three different conditions under which the same AI tool could scaffold or substitute self-regulation.

AI Tools in EFL Learning Contexts

For this review, AI is operationalized as any computational system that processes learner input, generates context-sensitive output, and, in more advanced implementations, adapts responses based on prior interaction (Roll & Wylie, 2016). Within EFL learning, the 22 included studies document AI tools across four functionally distinct categories, each interacting with SRL in theoretically different ways. Automated writing evaluation systems, specifically Grammarly and Write & Improve, provide immediate corrective feedback that can either stimulate self-evaluation at Zimmerman's self-reflection phase or suppress it, depending on whether learners engage actively or accept feedback passively (Lee & Coniam, 2013; Ranalli, 2018; Rogti & Ouarniki, 2026). Generative AI platforms such as ChatGPT, Poe, Bard, Tongyi.ai, ERNIE Bot, Kimi, and SpeakGuru (Kohnke et al., 2023) offer open-ended interaction that scaffolds brainstorming, drafting, and revision, though they also risk enabling learners to offload cognitive regulation, a substitution dynamic that Hu and Zhang (2025) and Tran et al. (2025) document directly. Adaptive reading platforms built on ChatGPT infrastructure adjust support intensity based on learner progress, primarily activating Plan-Enact-Reflect SRL phases (Pan et al., 2025, 2026). Specialized tools for speaking and vocabulary development, including ELSA Speak, Copilot via WhatsApp, and AI-VR agents, extend AI's SRL-relevant functions into skill domains the writing-dominated literature has largely underexplored (Tao et al., 2026; Teng, 2021).

AI and SRL: Evidence and Gaps

Research consistently demonstrates that AI-generated feedback activates the goal-setting and self-monitoring dimensions of SRL, particularly in writing contexts (Fitriati & Willian, 2025; Rogti & Ouarniki, 2026). The mechanisms through which this occurs are generally interpretable within Zimmerman's cyclical framework. When AI provides immediate, differentiated performance feedback, learners are prompted to revise goals and recalibrate strategies at the self-reflection phase. The motivational dimension of SRL presents a more contested picture. While several studies report enhanced intrinsic motivation following AI integration (Behforouz & Al Ghaithi, 2025; Chen & Alibakhshi, 2026), others document reduced autonomous motivation among high-AI-use learners (Tran et al., 2025), and Hu & Zhang (2025) identify a disengaged learner profile whose frequent AI use was decoupled from strategic self-regulation. The distinction between scaffolding and substitution remains undertheorized in the existing literature. Vygotskian scaffolding theory (Wood et al., 1976) implies that effective support gradually transfers cognitive control to the learner. AI tools that provide open-ended, on-demand assistance may invert this logic if learners develop tool dependency before internalizing regulatory strategies. This scaffolding versus substitution tension is the central unresolved question the present review addresses.

Method

Research Design

This study follows the Preferred Reporting Items for Systematic Reviews and Meta-Analyses (PRISMA 2020) framework (Page et al., 2021). A systematic design was selected because the research questions require explicit documentation of search strategies, selection procedures, and synthesis logic, conditions that narrative reviews cannot adequately guarantee (Snyder, 2019). Quality appraisal using the MMAT was integrated as a mandatory component, enabling evidence weighting during the synthesis phase (Hong et al., 2018). The protocol for this review was not pre-registered prior to data collection, which is acknowledged as a limitation. The search strategy, inclusion criteria, and appraisal procedures were nonetheless operationalized in writing and documented before screening commenced. Narrative synthesis with explicit quality weighting was adopted as the more appropriate analytical approach (Popay et al., 2006). The weighting mechanism applied a three-tier decision rule: studies scoring 4-5/5 on the MMAT anchored interpretive claims and were treated as primary evidence, studies scoring 3/5 were used to illustrate thematic patterns without supporting causal language, and no study scoring below 3/5 was retained in the final corpus. Inter-rater reliability for both the screening and quality appraisal phases is reported in full under Screening and Selection Procedures below (Cohen's kappa = .83 at the abstract stage and .79 at the full-text stage).

Search Strategy

Two major databases were searched, Scopus and Web of Science (WoS). The search was conducted in April 2026 using structured query strings. In Scopus: (“artificial intelligence” OR “AI” OR “chatbot” OR “generative AI” OR “ChatGPT”) AND (“self-regulated learning” OR “self-regulation” OR “learner autonomy”) AND (“EFL” OR “ESL” OR “English language learning”) AND PUBYEAR > 2019. In WoS: TS=(“artificial intelligence” OR “chatbot” OR “generative AI”) AND TS=(“self-regulated learning” OR “self-regulation”) AND TS=(“EFL” OR “ESL” OR “language learning”) AND PY=2020-2026. The initial search returned 178 records from Scopus and 133 from Web of Science, totaling 311 records. Following deduplication, 40 overlapping records were removed, yielding 271 unique records for title and abstract screening.

Inclusion and Exclusion Criteria

Seven criteria governed record selection, operationalized prior to screening to minimize selection bias. The SRL measurement criterion was applied to accommodate the theoretical plurality documented in the corpus. Studies were included if SRL was assessed using any validated instrument derived from an established framework, including Zimmerman’s three-phase model, Pintrich’s MSLQ, Izadpanah’s SDL scale, SDT-based need satisfaction measures, or equivalent scales with documented psychometric properties covering reported reliability coefficients and construct validity evidence. Studies employing single-item indicators or instruments without any reported psychometric evidence were excluded. This review does not anchor its synthesis to a single operational definition of SRL. Construct-level inclusiveness serves as the operational stance instead, treating SRL as any psychologically validated self-regulatory construct rather than compliance with one theoretical model. The corpus’s genuine plurality demands this choice, forcing 22 studies into one framework would discard the very diversity this review sets out to map. All seven inclusion criteria, research theme, population, research design, document type, language, publication year, and SRL measurement, are fully operationalized in Table 1 below, with explicit inclusion and exclusion conditions specified for each.

Table 1. Inclusion and Exclusion Criteria

Criterion	Inclusion	Exclusion
Research theme	Explicitly addresses AI in relation to SRL in EFL/ESL contexts	No explicit AI-SRL connection; SRL mentioned descriptively only
Population	EFL/ESL learners at any educational level	Native English speakers; non-language learning contexts
Research design	Empirical studies: qualitative, quantitative, or mixed methods	Conceptual papers, reviews, meta-analyses, opinion pieces
Document type	Peer-reviewed journal articles	Conference papers, theses, book chapters, grey literature

Criterion	Inclusion	Exclusion
Language	Published in English	Articles in languages other than English
Publication year	2020-2026	Before 2020
SRL measurement	SRL explicitly measured using a validated instrument derived from established frameworks (e.g., MSLQ, Zimmerman's model, Izadpanah's SDL, SDT-based scales) with reported reliability and construct validity evidence	SRL mentioned descriptively or assessed through single-item or unvalidated instruments without reported psychometric properties

Screening and Selection Procedures

Screening followed the four-stage PRISMA 2020 procedure illustrated in Figure 1. Of the 271 deduplicated records, abstract-level screening excluded 206 records, 141 lacked an explicit AI-SRL connection, 39 were conducted outside EFL/ESL contexts, and 26 did not meet document-type criteria. The remaining 65 records were retrieved for full-text assessment. Full-text screening excluded a further 43 records for the following reasons, SRL not formally measured with a validated instrument ($n = 12$), non-EFL/ESL context ($n = 7$), no empirical data reported ($n = 6$), publication language or scope outside the inclusion criteria ($n = 5$), insufficient peer-review documentation ($n = 8$), and AI tool not clearly operationalized in relation to SRL outcomes ($n = 5$). The final corpus comprised 22 articles.

To establish inter-rater reliability, a randomly selected 20% subsample ($n = 54$ of 271 deduplicated records) was independently screened by a second reviewer at the abstract stage. Cohen's kappa was $k = .83$, indicating strong agreement. At the full-text stage, all 65 records were reviewed independently by both reviewers, yielding $k = .79$. The most frequent source of disagreement at both stages concerned the SRL measurement criterion, particularly for studies using hybrid or adapted instruments. Disagreements were resolved through discussion and conservative application of the inclusion criterion. When adequate psychometric documentation was absent, the record was excluded.

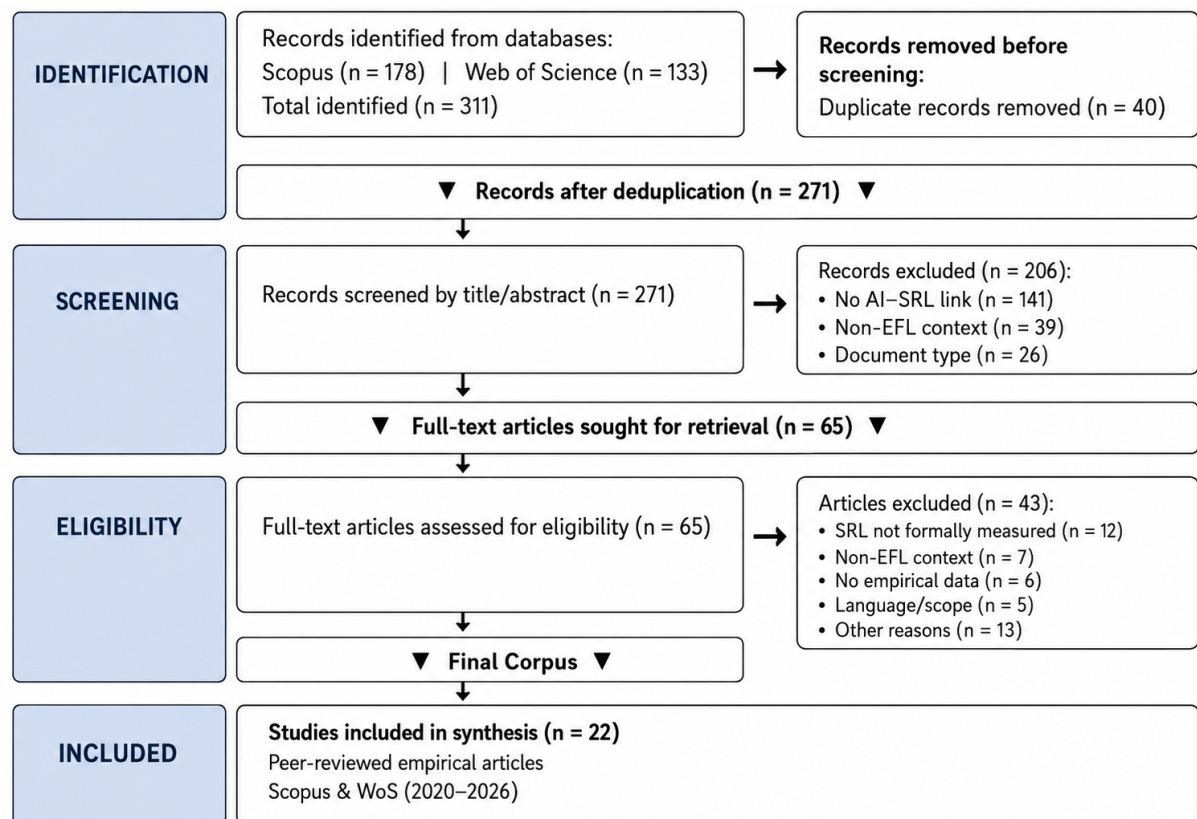


Figure 1. PRISMA 2020 Flow Diagram of Study Selection

Quality Appraisal

All 22 included studies were assessed using the MMAT (Hong et al., 2018), a validated tool designed to appraise qualitative, quantitative, and mixed-methods studies within a single framework. Fourteen studies (63.6%) achieved scores of 4-5/5, indicating low risk of bias. Six studies (27.3%) scored 3/5, reflecting moderate risk due primarily to small samples or inadequately reported sampling procedures. Two qualitative studies received 3/5 because sample sizes were insufficient to support analytic generalization. Quality differentials are reflected throughout the synthesis, high-quality studies anchor interpretive claims, while moderate-quality studies illustrate patterns rather than establish them.

Table 2. MMAT Quality Appraisal Summary Across Included Studies (N = 22)

MMAT Score	Studies (n, %)	Role in Synthesis
4-5/5 (Low risk of bias)	14 (63.6%)	Anchor interpretive claims, effect sizes weighted accordingly
3/5 (Moderate risk of bias)	6 (27.3%)	Illustrate thematic patterns, findings presented cautiously without causal claims
3/5 qualitative (Limited generalizability)	2 (9.1%)	Descriptive mention only, excluded from effect-size comparisons

Data Extraction and Synthesis

Data were extracted from each included article using a standardized protocol covering the following fields, author(s) and year, country and EFL context, educational level, AI tool type and implementation mode, SRL theoretical framework and measurement instrument, research design, sample size, key findings, reported effect sizes, and MMAT score. Synthesis followed a narrative-thematic approach (Popay et al., 2006) organized around the three research questions. Thematic categories were constructed deductively, anchored in the SRL frameworks represented in the corpus (Pintrich, 2004; Zimmerman, 2000), then refined inductively as recurring patterns emerged across studies. RQ1 findings were synthesized through frequency analysis and descriptive mapping of AI tool types and theoretical frameworks. RQ2 was addressed through cross-study thematic comparison of SRL dimensions, with effect directions coded systematically as positive, null, or negative for each dimension reported. RQ3 findings were derived from pattern-based moderator analysis. Coding categories were verified by the second reviewer on a 30% subsample and disagreements were resolved through discussion.

Results and Discussions

Table 3 presents an overview of the included studies with key characteristics relevant to the three research questions addressed in this review. The variation in MMAT scores across studies, ranging from 3/5 to 5/5, reflects differences in sampling rigor rather than differences in tool efficacy. The six studies scoring 3/5 share small or unreported sample sizes (Al-smadi et al., 2025; Fitriati & Williyan, 2025; Gu & Liu, 2025; Jalambo et al., 2025; Qiao & Zhao, 2023; Yifan et al., 2026), a methodological limitation independent of which AI tool or SRL framework was used. This distinction matters for interpretation, since lower MMAT scores caution against causal claims but do not by themselves indicate weaker AI-SRL relationships.

Table 3. Overview of Included Studies (N = 22)

Author(s) & Year	Country	AI Tool	SRL Framework	Design	N	MMA T	Key SRL Outcome
Aladini et al. (2025)	Iran	AI-based writing tasks	Izadpanah SDL	Experimenta 1	561	4/5	Improved L2 writing, growth mindfulness, grammatical knowledge

Author(s) & Year	Country	AI Tool	SRL Framework	Design	N	MMA T	Key SRL Outcome
Al-Smadi et al. (2025)	Jordan	ChatGPT	SRL (general)	Focus Group	~25	3/5	Four learner profiles in AI- supported autonomous learning
Chen & Alibakhsh i (2026)	Iran	ChatGPT, Poe, Bard	Pintrich MSLQ	Experimenta 1	310	5/5	AI tools enhanced SRL strategy use and language achievement
Gu & Liu (2025)	China	AI chatbot	AGT (Chazan)	Mixed- methods	61	3/5	Goal-setting improved; over-reliance concerns documented
Hongxia & Razali (2026)	Malaysia	DeepSeek AI	Autonomous learning / self-efficacy	Mixed methods	N/ R	4/5	AI-assisted peer feedback improved SRL and motivation
Hu & Zhang (2025)	China	Generative AI	CVT (Pekrun)	Latent Profile (LPA)	551	5/5	Three SRL profiles: Proactive, Baseline, Disengaged
Jalambo et al. (2025)	Oman	ChatGPT	Self-reg. vocabulary learning	Quasi- experimental	N/ R	3/5	ChatGPT facilitated vocabulary SRL; reduced L2 boredom
Pan et al. (2025)	China	ReadMate (GPT 3.5)	Plan-Enact- Reflect / Zimmerman	Quasi- experimental	61	4/5	Improved self- regulated reading strategy use and engagement
Pan et al. (2026)	China	ReadMate (GenAI)	SDL & SRL (Zimmerman)	Design- Based Research	116	5/5	Interest-based GenAI showed greatest potential for SRL reading
Qiao & Zhao (2023)	China	AI-based tools	Self- regulation (general)	Descriptive	N/ R	3/5	AI-based learning positively impacted speaking and SRL

Author(s) & Year	Country	AI Tool	SRL Framework	Design	N	MMA T	Key SRL Outcome
Rogti & Ouarniki (2026)	Algeria	Grammarly , Write & Improve	Metacognitive scaffolding (Zimmerman)	Quantitative experimental	120	4/5	AI metacognitive scaffolding strengthened SRL and writing
Tao et al. (2026)	China	GenAI iVR agent	Zimmerman SRL phases	Randomized experimental	71	5/5	GenAI-VR triggered more frequent reflection and SRL strategies
Tran et al. (2025)	Vietnam / Australia	MT & GenAI (ChatGPT)	Self-Determination Theory	Mixed-effects modelling	N/R	4/5	Autonomy and competence mediated by disciplinary context
Wang (2024)	China	ChatGPT, DeepL, MT	Cognitive-sociocultural (Vygotsky)	Semester-long journal	79	4/5	Training improved ethical AI awareness and SRL strategy use
Wei (2023)	China	AI platform	SRL & L2 Motivation (Zimmerman)	Experimental	60	4/5	AI instruction improved achievement, motivation, and SRL
Xun et al. (2025)	China	AI-SREPS model	SRL & Positive Psychology (Zimmerman)	Mixed-design repeated	88	4/5	AI-SREPS improved speech quality and emotional engagement
Yifan et al. (2026)	China	Adaptive AI systems	Metacognitive learning (Teng)	Qualitative (interviews)	8	3/5	High digital literacy enabled cognitive amplification via AI
Zhang (2025)	China	ERNIE Bot, Kimi, SpeakGuru	SRL, Resilience & Autonomy (hybrid)	Quasi-experimental	183	4/5	AI improved SRL, resilience, and learner autonomy
Behforouz & Al Ghaithi (2025)	Oman	Copilot via WhatsApp	SRL (Zimmerman) & SDT	Random experimental	60	4/5	Copilot enhanced vocabulary mastery and SRL goal-setting

Author(s) & Year	Country	AI Tool	SRL Framework	Design	N	MMA T	Key SRL Outcome
Fitriati & Willian (2025)	Indonesia	ELSA, ChatGPT, Gemini	SRL phases (Zimmerman)	Qualitative case study	12	3/5	AI tools enhanced EFL presentations across all SRL phases
Rahimi et al. (2025)	Iran	ChatGPT	PESRL & PEL2MSS (PLS-SEM-ANN)	Hybrid PLS-SEM-ANN	133	5/5	ChatGPT shaped L2 motivation and SRL; digital authenticity key
Wang et al. (2025)	China	Tongyi.ai (GenAI)	Task strategies & Metacognition (Pintrich)	Mixed-methods quasi-exp.	40	4/5	GAI enabled adaptive task strategies and metacognitive externalization

RQ1: AI Tool Types and SRL Theoretical Frameworks

Across the 22 included studies, AI tools cluster into four functional categories. Generative AI platforms, including ChatGPT, Poe, Bard, Tongyi.ai, ERNIE Bot, Kimi, and SpeakGuru, constituted the largest category ($n = 14$, 64%), reflecting the rapid diffusion of large language model-based tools within the 2020-2026 publication window. Automated writing evaluation systems appeared in four studies (18%), adaptive reading platforms in two (9%), and specialized AI tools for speaking and vocabulary in two additional studies (9%).

Table 4. Distribution of AI Tool Categories Across Included Studies ($N = 22$)

AI Tool Category	n	%	Representative Studies
Generative AI (ChatGPT and derivatives: ERNIE Bot, Kimi, Tongyi.ai, SpeakGuru, Poe, Bard)	14	64%	Rahimi et al. (2025); Chen & Alibakhshi (2026); Wang et al. (2025); Tran et al. (2025); Gu & Liu (2025); Zhang (2025)
Automated Writing Evaluation / AWE (Grammarly, Write & Improve)	4	18%	Rogti & Ouarniki (2026); Fitriati & Willian (2025); Aladini et al. (2025); Wang (2024)
Adaptive AI Reading Platforms (ReadMate / ChatGPT 3.5)	2	9%	Pan et al. (2025); Pan et al. (2026)
Specialized AI for Speaking/Vocabulary (ELSA Speak, Copilot via WhatsApp, GenAI-iVR)	2	9%	Behforouz & Al Ghaithi (2025); Tao et al. (2026); Xun et al. (2025)

The dominance of generative AI platforms reflects a qualitative shift in the AI-SRL relationship that earlier literature did not anticipate. AWE systems operate within a bounded corrective feedback loop, the learner submits text, receives error-specific feedback, and decides whether to revise. Generative AI platforms introduce an open-ended dialogic dynamic across multiple turns, a structure that potentially activates all three phases of Zimmerman's cycle simultaneously but also creates conditions for regulatory substitution that AWE systems do not. Hu and Zhang (2025) latent profile analysis captures this tension precisely.

Frequent AI use did not uniformly predict high SRL quality, a disengaged profile characterized learners who engaged with AI tools extensively but without strategic intentionality. China was the most represented context ($n = 12$, 55%), followed by Iran ($n = 3$), Oman ($n = 2$), and single-country studies from Indonesia, Malaysia, Algeria, Vietnam/Australia, and Jordan. This concentration introduces a potential confound, given that Chinese EFL learners operate within gaokao-oriented, teacher-directed examination cultures that may systematically shape how AI tools interact with self-regulatory behavior. By language skill domain, writing was the most frequently studied ($n = 10$, 45%), followed by vocabulary and speaking ($n = 4$ each, 18%) and reading ($n = 3$, 14%). Vocabulary studies, notably, Behforouz and Al Ghaithi (2025) and Jalambo et al. (2025) consistently documented goal-setting and monitoring gains alongside reduced L2 boredom.

Table 5. Geographic Distribution of Included Studies ($N = 22$)

Country/Region	n	Representative Studies
China	12	Hu & Zhang (2025); Pan et al. (2025, 2026); Zhang (2025); Xun et al. (2025); Gu & Liu (2025); Wang et al. (2025); Tao et al. (2026); Wang (2024); Qiao & Zhao (2023); Yifan et al. (2026); Wei (2023)
Iran	3	Rahimi et al. (2025); Aladini et al. (2025); Chen & Alibakhshi (2026)
Oman	2	Behforouz & Al Ghaithi (2025); Jalambo et al. (2025)
Indonesia	1	Fitriati & Williyani (2025)
Malaysia	1	Hongxia & Razali (2026)
Algeria	1	Rogti & Ouarniki (2026)
Vietnam / Australia	1	Tran et al. (2025)
Jordan	1	Al-Smadi et al. (2025)

The SRL theoretical landscape shows considerable pluralism rather than the nominal convergence that citation counts might suggest. Zimmerman's cyclical model appeared in nine studies (41%), Pintrich's MSLQ-based framework in four (18%), and SDT in three (14%). Izadpanah's SDL, Control-Value Theory, and Achievement Goal Theory each appeared in one study, and three studies employed hybrid frameworks (Table 6).

Table 6. SRL Theoretical Frameworks Across Included Studies ($N = 22$)

SRL Framework	n	%	Representative Studies
Zimmerman's Cyclical SRL Model (2000/2002)	9	41%	Pan et al. (2025, 2026); Tao et al. (2026); Xun et al. (2025); Rogti & Ouarniki (2026); Behforouz & Al Ghaithi (2025); Fitriati & Williyani (2025)
Pintrich's MSLQ / Cognitive-Metacognitive Framework	4	18%	Chen & Alibakhshi (2026); Wang et al. (2025); Wei (2023); Zhang (2025)
Self-Determination Theory / SDT (Ryan & Deci, 2000)	3	14%	Tran et al. (2025); Behforouz & Al Ghaithi (2025); Rahimi et al. (2025)
Other/Hybrid (PESRL+PEL2MSS, Positive Psychology, Resilience)	3	14%	Rahimi et al. (2025); Xun et al. (2025); Zhang (2025)
Izadpanah's SDL Model	1	5%	Aladini et al. (2025)
Control-Value Theory / CVT (Pekrun, 2006)	1	5%	Hu & Zhang (2025)
Achievement Goal Theory / AGT (Chazan & Daniel, 2022)	1	5%	Gu & Liu (2025)

RQ2: AI Effects on SRL Dimensions

The SRL dimensions most consistently affected by AI integration were goal-setting, self-monitoring and evaluation, and metacognitive strategy use. Goal-setting improvements were reported in 15 of the 22 studies (68%), frequently appearing as the most prominent outcome. Tools providing immediate, differentiated performance feedback create conditions that directly prompt goal revision and recalibration, which is the core mechanism of Zimmerman's forethought and self-reflection phases. Behforouz and Al Ghaithi (2025) and Chen and Alibakhshi (2026) both document goal-setting as the most robust SRL gain across their respective interventions.

Table 7. Summary of AI Effects on SRL Dimensions Across Included Studies (N = 22)

SRL Dimension	Positive Effect (n, %)	Null/Negative (n, %)	Key Evidence
Goal-setting	15 (68%)	7 (32%)	Behforouz & Al Ghaithi (2025); Chen & Alibakhshi (2026); Pan et al. (2025)
Self-monitoring / self-evaluation	13 (59%)	9 (41%)	Rogti & Ouarniki (2026); Fitriati & Williyani (2025); Wang et al. (2025)
Metacognitive strategy use	11 (50%)	11 (50%)	Zhang (2025); Xun et al. (2025); Pan et al. (2026)
Motivational regulation (autonomous)	10 (45%)	12 (55%)	Tran et al. (2025) -- negative; Hu & Zhang (2025) -- Disengaged profile
Effort regulation / persistence	8 (36%)	14 (64%)	Gu & Liu (2025); Tran et al. (2025)
Resource management	6 (27%)	16 (73%)	Wei (2023); Wang (2024)

Self-monitoring gains appeared in 13 studies (59%) and metacognitive awareness in 11 (50%). These results are consistent across tool categories. Even AWE systems, whose feedback is relatively bounded, activated self-monitoring in writing contexts (Rogti & Ouarniki, 2026). The more theoretically significant finding concerns motivational regulation. Positive motivational effects appeared in 10 studies, while null or negative effects emerged in 12. Most prominently, Hu and Zhang (2025) identified a Disengaged SRL profile among frequent AI users, and Tran et al. (2025) documented reduced autonomous motivation among learners reporting high AI dependence.

Measurement practices across the corpus also warrant scrutiny. Seventeen of the 22 studies (77%) relied exclusively on self-report questionnaires, including Qiao and Zhao (2023) whose descriptive design, while providing useful baseline insights on AI-integrated speaking and self-regulation, could not establish directional mechanisms, capturing learners' perceptions of their regulation rather than behavioral evidence of it. Only five studies incorporated process-level data, including log files (Pan et al., 2025, 2026; P. Wang et al., 2025), VR-based sequence analysis (Tao et al., 2026), and reflective journals (Wang, 2024). These five produced the most mechanistic accounts of AI-SRL interaction, which suggests that the field's predominant measurement choices impose a substantive constraint on what it can claim to know.

RQ3: Moderating Factors

Among the moderating factors identified across the corpus, instructional integration quality emerges as the most consistent. Studies that embedded AI within structured pedagogical sequences, including the AI-assisted peer feedback workflow in Hongxia and Razali (2026), the six-phase AI-SREPS protocol in Xun et al. (2025), and the Plan-Enact-Reflect chatbot scaffolding in Pan et al. (2025), consistently produced larger SRL gains than studies providing unrestricted AI access. This pattern aligns with Vygotskian scaffolding theory (Wood et al., 1976), AI support is most effective when designed to transfer cognitive control to the learner over time.

Digital literacy, however, also shaped outcomes in ways that cut across tool categories. Yifan et al. (2026) found that high digital literacy correlated with cognitive amplification, a measurable shift in metacognitive approach, while low literacy produced frustration and SRL disengagement. This direction is consistent with Zhang (2025) larger quasi-experimental study (n = 183), which documented differential autonomy gains across learners with varying prior technology experience. Intervention duration also mattered. Studies running eight weeks or longer (n = 9) showed more consistent gains across multiple SRL dimensions; shorter

interventions typically produced improvements in goal-setting but without corresponding shifts in metacognitive regulation or motivational autonomy.

Over-reliance risk is perhaps the most consequential finding for practitioners. Across high-quality studies, this emerges not as an isolated occurrence but as a recurring structural pattern. Gu and Liu (2025) explicitly documented over-reliance concerns in a chatbot-supported EFL context. Tran et al. (2025) found reduced autonomous motivation among high-AI-use learners. Hu and Zhang (2025) Disengaged profile, representing a substantial subgroup of a sample of 551 learners, describes learners whose frequent AI use was decoupled from strategic self-regulation. This cross-contextual pattern, appearing in both AWE and generative AI environments, suggests that what shapes learner outcomes is not the tool per se, but the instructional design surrounding its deployment. Whether AI functions as a cognitive scaffold or a substitution mechanism depends on how tasks are structured, monitored, and debriefed. Skill domain also moderated effect magnitude. Writing-focused interventions produced the largest SRL gains in this corpus. Chen and Alibakhshi (2026), whose three-tool experimental comparison involved 310 participants, yielded the highest strategy gains across all studies reviewed. Aladini et al. (2025) reached a comparable conclusion through a large-scale writing intervention ($n = 561$), linking AI-integrated task completion to improvements in grammatical knowledge and growth mindfulness alongside measurable SRL development.

Discussion

The synthesis confirms that AI integration in EFL learning supports self-regulated learning, though the robustness of that support depends on which SRL component is examined, which tool is used, how SRL is measured, and what instructional conditions frame AI deployment. Hu and Zhang's (2025) latent profile analysis represents the most theoretically significant contribution in this corpus on three explicit grounds. It is the only study to empirically derive learner subgroups rather than assume homogeneity within the sample, it directly tests rather than presumes the scaffolding-versus-substitution distinction that the present review treats as its central theoretical tension, and its 551-participant sample exceeds the median sample size of the corpus by a wide margin, supporting greater confidence in the profile boundaries it identifies. By comparison, frameworks such as Pintrich's MSLQ and Ryan and Deci's SDT, while theoretically informative, were applied in this corpus through variance-based or correlational designs that test pre-specified relationships rather than discover them inductively from the data. This is a different inferential contribution, not a lesser one, but it does not address the disengaged-versus-proactive distinction directly. AI use does not map linearly onto SRL quality, and a substantial subgroup of active AI users may develop disengaged rather than proactive self-regulatory profiles. This finding is in line with Al-Smadi et al. (2025), whose qualitative profiling of EFL learners in Jordan identified four distinct response patterns to AI-supported autonomous learning, ranging from highly engaged to avoidant.

The consistent positive effects on goal-setting and self-monitoring are interpretable through Zimmerman's cyclical model. AI-generated feedback creates frequent forethought-performance-reflection micro-cycles that prompt goal revision and monitoring. The stronger effects in writing contexts reflect writing's natural alignment with cyclical revision processes. Motivational regulation remains the most contested dimension in the corpus. Whether AI-supported regulation is genuinely internalized or simply enacted in response to tool affordances is a question self-report measures cannot answer. SDT grounded designs with behavioral components are uniquely equipped to address this distinction (Ryan & Deci, 2000), though hybrid analytical approaches, as in Rahimi et al. (2025) PLS-SEM-ANN model, suggest that digital authenticity and perceived credibility of AI feedback are mediating variables that warrant more systematic treatment in future motivational designs.

A persistent limitation of the corpus is insufficient theorization of the EFL context as a moderating variable. Most studies treat EFL as background, a descriptor of the population rather than a theoretically meaningful condition that interacts with specific AI affordances to produce distinct SRL patterns. Indonesia, represented by a single qualitative study (Fitriati & Williyani, 2025), and Malaysia, by one mixed-methods study (Hongxia & Razali, 2026), signal the existence of Southeast Asian EFL-AI-SRL research but cannot yet support context-specific conclusions. Operationally, EFL context functions as a moderating rather than descriptive variable when at least one of three conditions varies systematically across the corpus: amount of authentic L2 input available outside instruction, the examination orientation of the educational system, and the degree of teacher-directiveness in classroom culture; future coding schemes should record these three indicators explicitly for each study rather than treating country name as a proxy for context. Similarly, "instructional conditions" in this review refers concretely to three observable features, whether AI use was embedded in a structured task sequence with defined phases, whether human feedback supplemented AI feedback, and whether learners received explicit instruction in regulating their own AI use, each of which a

future coder could apply with a target inter-rater kappa above 0.80. For Indonesian EFL practitioners, the implication is concrete, embed AI within scaffolded sequences with deliberate, time-bound reduction of support, and treat digital literacy as a prerequisite rather than an assumption. Developing this research infrastructure in high-population EFL environments, including Indonesia, Brazil, and sub-Saharan Africa, represents the field's most consequential geographic priority.

When studies employ six different SRL frameworks, each operationalized through distinct instruments with varying psychometric documentation, nominal comparisons across studies rest on a fragile epistemic base. The five studies that operationalized SRL through behavioral trace data produced higher MMAT scores and more mechanistic accounts of AI-SRL interaction. The field would benefit from agreed-upon reporting standards for intervention design transparency, analogous to what PRISMA 2020 provides for synthesis methodology.

Conclusions

Drawing on 22 empirical studies published between 2020 and 2026, and following PRISMA 2020 protocol with MMAT quality appraisal, this review yields three substantive conclusions. First, the AI tool landscape has shifted decisively toward generative AI platforms. That shift is not merely quantitative, it reframes fundamental questions about the AI-autonomy relationship in ways earlier AWE-centered research did not anticipate. Theoretically, the corpus is more pluralistic than its surface Zimmerman-dominance suggests, SDL, CVT, AGT, SDT, and multiple hybrid configurations each appear, reflecting the genuine complexity of SRL as a construct. This diversity is best understood as a methodological liability layered on top of a substantive strength. The breadth of frameworks captures real facets of self-regulation that any single model would flatten, yet because each framework operationalizes SRL through non-interchangeable instruments, the corpus cannot currently support meta-analytic effect-size pooling across studies, only narrative comparison of directional patterns. Future syntheses would benefit from a published framework-to-instrument crosswalk, of the kind summarized in Table 6 of this review, so that researchers can identify which findings are directly comparable and which are not. Second, AI integration most reliably supports goal-setting, self-monitoring, and metacognitive strategy use. Its effects on motivational regulation are considerably more context- and tool-dependent. The most theoretically consequential finding across the corpus is Hu and Zhang (2025) demonstration that frequent AI use does not uniformly predict high SRL quality, reframing the research agenda from whether to how. Third, measurement inconsistency across six SRL frameworks, multiple instruments, and predominantly self-report designs limits the cumulative evidentiary value of this body of literature. The five studies that incorporated behavioral trace data produced the most credible and theoretically precise accounts of AI-SRL interaction. This review set the minimum acceptable MMAT score for inclusion at 3/5. Scores below this threshold typically signal unreported sampling procedures or an insufficient justification of analytic choices, and under such conditions even directional findings cannot be interpreted with confidence. A higher cutoff was deliberately avoided, however, since it would have excluded qualitative and small-sample studies whose contextual richness, particularly from underrepresented EFL settings, would otherwise be lost from the synthesis. For researchers, the most productive directions involve behavioral trace data alongside self-report, complete operationalizations of single frameworks, and longitudinal designs capable of tracking SRL trajectory changes over extended periods. For practitioners, the evidence supports embedding AI within explicitly scaffolded pedagogical sequences with planned reduction of AI support over time, and attending to digital literacy as a precondition rather than an assumed competency. Indonesian secondary and tertiary EFL settings, characterized by limited digital literacy infrastructure and exam-driven motivational climates, require preparation-phase interventions before AI-integrated SRL programs can be expected to yield consistent benefits. Limitations of this review include restriction to Scopus and WoS, exclusion of non-English publications, absence of pre-registration, and five studies not reporting sample sizes explicitly. Future research should address these gaps through multilingual database searches, collaborative international designs, and mixed-methods approaches integrating behavioral trace data with qualitative perspectives on learner experience.

Acknowledgments

The author gratefully acknowledges the guidance and scholarly support of supervisors and colleagues throughout this research process.

References

- Al-Smadi, O. A., Ab, R., & Al-Ramahi, R. A. (2025). Jordanian English Language Learners' Engagement With Ai-Supported Self-Regulated Learning. *Research In Learning Technology*, 33(1063519), 1–18.
- Aladini, A., Ismail, S. M., Ahmad, M., Khasawneh, S., & Shakibaei, G. (2025). Computers In Human Behavior Reports Self-Directed Writing Development Across Computer / Ai-Based Tasks : Unraveling The Traces On L2 Writing Outcomes , Growth Mindfulness , And Grammatical Knowledge. *Computers In Human Behavior Reports*, 17(November 2024), 100566. <https://doi.org/10.1016/j.chbr.2024.100566>
- Behforouz, & Ghaithi, A. (2025). Ai As A Language Learning Facilitator : Examining Vocabulary And Self-Regulation In Efl Learners. *International Journal Of Learning, Teaching And Educational Research*, 16, 616–634.
- Broadbent, J., & Poon, W. L. (2015). Self-Regulated Learning Strategies And Academic Achievement In Online Higher Education Learning Environments: A Systematic Review. *Internet And Higher Education*, 27, 1–13. <https://doi.org/10.1016/j.iheduc.2015.04.007>
- Chazan, D. J., Pelletier, G. N., & Daniels, L. M. (2022). *Achievement Goal Theory Review : An Application To School Psychology*. <https://doi.org/10.1177/08295735211058319>
- Chen, J., & Alibakhshi, G. (2026). Powered Applications' Effects On English Language Learners' Cognitive , Metacognitive , And Resource Management Strategies , And Language Achievement. *Journal Of Computer Assisted Learning*, 41. <https://doi.org/10.1002/jcal.70171>
- Chiu, T. K. F. (2024). Computers And Education : Artificial Intelligence Future Research Recommendations For Transforming Higher Education With Generative Ai. *Computers And Education: Artificial Intelligence*, 6(December 2023), 100197. <https://doi.org/10.1016/j.caeai.2023.100197>
- Fitriati, S. W., & Williyani, A. (2025). *Ai-Enhanced Self-Regulated Learning : Efl Learners' Prioritization And Utilization In Presentation Skills Development*. 9(2), 22–37.
- Gu, J., & Liu, Q. (2025). Enhancing Chinese Efl University Students' Self-Regulated Learning Through Ai Chatbot Intervention : Insights From Achievement Goal Theory. *Learning And Motivation*, 92(September), 102191. <https://doi.org/10.1016/j.lmot.2025.102191>
- Habók, A., & Magyar, A. (2022). *English As A Foreign Language Learners' Strategy Awareness Across Proficiency Levels From The Perspective Of Self-Regulated Learning Metafactors*. October, 1–13. <https://doi.org/10.3389/fpsyg.2022.1019561>
- Hong, Q. N., Pluye, P., Fàbregues, S., Bartlett, G., Boardman, F., Cargo, M., Dagenais, P., Pierre, M., Griffiths, F., Nicolau, B., Cathain, A. O., Claude, M., & Vedel, I. (2018). Mixed Methods Appraisal Tool (Mmat) Version 2018 User Guide. *Canadian Intellectual Property Office*.
- Hongxia, H., & Razali, A. B. (2026). Impact Of Ai- Assisted Peer Feedback On Efl Students' Self -Regulated. *International Journal Of Instruction*, 19(2), 163–182.
- Hu, J., & Zhang, H. (2025). Exploring Ai - Assisted Self - Regulated Learning Profiles And The Predictive Role Of Academic Appraisals : A Control- - Value Perspective On Chinese Efl University Students. *European Journal Of Education*, 2024, 1–10. <https://doi.org/10.1111/Ejed.70249>
- Ismail, S. M., Nikpoo, I., & Prasad, K. D. V. (2023). Promoting Self - Regulated Learning , Autonomy , And Self - Efficacy Of Efl Learners Through Authentic Assessment In Efl Classrooms. *Language Testing In Asia*. <https://doi.org/10.1186/S40468-023-00239-Z>
- Izadpanah, S. (2022). *The Impact Of Flipped Teaching On Efl Students' Academic Resilience , Self-Directed Learning , And Learners' Autonomy*. 05(December), 1–12. <https://doi.org/10.3389/fpsyg.2022.981844>
- Jalambo, M. O., Çakmak, F., & Akhter, S. (2025). Effects Of Self-Regulated Vocabulary Learning With Chatbots On Incidental And Collocational Vocabulary Learning And Foreign Language Learning Boredom. *Discover Education*.
- Kasneci, E., Sessler, K., Stefan, K., Bannert, M., Fischer, F., Gasser, U., Groh, G., Stephan, G., Poquet, O., Sailer, M., Schmidt, A., & Seidel, T. (2023). Chatgpt For Good? On Opportunities And Challenges Of Large Language Models For Education. *Learning And Individual Differences*, 1–13. <https://doi.org/10.1016/j.lindif.2023.102274>
- Kohnke, L., Moorhouse, B. L., & Zou, D. (2023). *Chatgpt For Language Teaching And Learning*. 1–14. <https://doi.org/10.1177/00336882231162868>
- Lee, I., & Coniam, D. (2013). Introducing Assessment For Learning For Efl Writing In An Assessment Of Learning Examination-Driven System In Hong Kong. *Journal Of Second Language Writing*, 22(1), 34–50. <https://doi.org/10.1016/j.jslw.2012.11.003>
- Page, M. J., Matthew, J., Joanne, E., Patrick, M., Mckenzie, J. E., Bossuyt, P. M., Boutron, I., Hoffmann,

- T. C., Mulrow, C. D., Shamseer, L., Tetzlaff, J. M., Akl, E. A., Brennan, S. E., Chou, R., Glanville, J., Grimshaw, J. M., Hróbjartsson, A., Lalu, M. M., Li, T., ... Welch, V. A. (2021). *The Prisma 2020 Statement: An Updated Guideline For Reporting Systematic Reviews*. <https://doi.org/10.1136/bmj.N71>
- Pan, M., Lai, C., & Guo, K. (2025). Effects Of Genai-Empowered Interactive Support On University Efl Students ' Self-Regulated Strategy Use And Engagement In Reading. *The Internet And Higher Education*, 65(October 2024).
- Pan, M., Lai, C., Guo, K., & Huang, J. (2026). Computers & Education Self-Regulation Plus Individual Interests ? A Design-Based Study On The Development Of A Genai-Empowered Platform For Self-Directed. *Computers & Education*, 241(April 2025), 105474. <https://doi.org/10.1016/j.compedu.2025.105474>
- Pekrun, R. (2024). Control - Value Theory : From Achievement Emotion To A General Theory Of Human Emotions. In *Educational Psychology Review*. Springer Us. <https://doi.org/10.1007/S10648-024-09909-7>
- Pintrich, P. R. (2004). A Conceptual Framework For Assessing Motivation And Self-Regulated Learning In College Students. *Educational Psychology Review*, 16(4), 385–407.
- Popay, J., Roberts, H., Petticrew, M., Roen, K., & Duffy, S. (2006). Guidance On The Conduct Of Narrative Synthesis In Systematic Reviews A Product From The Esrc Methods Programme With. *Esrc Research Methods Programme*, April 2006, 1–92.
- Qiao, H., & Zhao, A. (2023). Artificial Intelligence-Based Language Learning : Illuminating The Impact On Speaking Skills And Self-Regulation In Chinese Efl Context. *Frontiers In Psychology*, November. <https://doi.org/10.3389/fpsyg.2023.1255594>
- Rahimi, R. A., Sheykhholeslami, M., & Pour, A. M. (2025). Uncovering Personalized L2 Motivation And Self-Regulation In Chatgpt-Assisted Language Learning: A Hybrid Pls-Sem-Ann Approach. *Computers In Human Behavior Reports*, 17(September 2024), 100539. <https://doi.org/10.1016/j.chbr.2024.100539>
- Ranalli, J. (2018). Automated Written Corrective Feedback : How Well Can Students Make Use Of It ? Students Make Use Of It ? *Elt Journal*, 72(3), 318–333. <https://doi.org/10.1080/09588221.2018.1428994>
- Rogti, M., & Ouarniki, L. (2026). From Error Correction To Strategy Support : Ai-Powered Metacognitive Scaffolding In Assessing Efl Writing Performance And Self-Regulatory Engagement. *Language Teaching Research*. <https://doi.org/10.1177/13621688261420086>
- Roll, I., & Wylie, R. (2016). *Evolution And Revolution In Artificial Intelligence In Education*. 582–599. <https://doi.org/10.1007/S40593-016-0110-3>
- Ryan, R. M., & Deci, E. L. (2000). *Self-Determination Theory And The Facilitation Of Intrinsic Motivation, Social Development, And Well-Being*. 55(1), 68–78.
- Snyder, H. (2019). Literature Review As A Research Methodology : An Overview And Guidelines. *Journal Of Business Research*, 104(August), 333–339. <https://doi.org/10.1016/j.jbusres.2019.07.039>
- Tao, L., Song, Y., & Fu, J. (2026). Computers & Education Exploring Students ' Self-Regulated Learning Behavioural Patterns And Perceptions In An English Speaking Task Within A Generative Ai-Supported Immersive Vr Environment. *Computers & Education*, 242(October 2025), 105515. <https://doi.org/10.1016/j.compedu.2025.105515>
- Teng, L. S. (2021). Individual Differences In Self-Regulated Learning : Exploring The Nexus Of Motivational Beliefs , Self- Efficacy , And Srl Strategies In Efl Writing. *Language Teaching Research*, 28(2), 366–388. <https://doi.org/10.1177/13621688211006881>
- Tran, H., Crosthwaite, P., Huynh, Q., & Pham, P. (2025). *Journal Of English For Academic Purposes Students ' Self-Determination In Using Machine Translation And Generative Ai Tools For English For Academic Purposes*. 78(September).
- Wang, P., Liu, T., Yang, Y., & Xiang, X. (2025). Optimizing Self- - Regulated Learning : A Methods Study On Gai ' S Impact On Undergraduate Task Strategies And Metacognition. *British Journal Of Educational Technology*, August, 1–20. <https://doi.org/10.1111/Bjet.70018>
- Wang, Y. (2024). Cognitive And Sociocultural Dynamics Of Self-Regulated Use Of Machine Translation And Generative Ai Tools In Academic Efl Writing. *System*, 126(October), 103505. <https://doi.org/10.1016/j.system.2024.103505>
- Wood, D., Bruner, J. S., & Ross, G. (1976). The Role Of Tutoring In Problem Solving. *Journal Of Child Psychology And Psychiatry And Allied Disciplines*, 17(2), 89–100. <https://doi.org/10.1111/J.1469-7610.1976.tb00381.x>
- Xun, H., Zhou, C., Li, Y., & Wu, Y. (2025). Ai-Supported English Public Speaking In The Chinese Efl Context : Insights From Self-Regulated Learning And Positive Psychology. *Learning And Motivation*,

- 92(September), 102211. <https://doi.org/10.1016/J.Lmot.2025.102211>
- Yifan, P., Hashim, H., & Said, N. E. M. (2026). Cognitive Amplification: Harnessing Artificial Intelligence To Augment Metacognitive Learning Strategies. *Dirasat: Human And Social Sciences*, 53(1), 1–13.
- Zawacki-Richter, O., Marin, V. I., & Bond, M. (2019). *Systematic Review Of Research On Artificial Intelligence Applications In Higher Education – Where Are The Educators ?*
- Zhang, Z. (2025). *The Role Of Artificial Intelligence Tools On Chinese Efl Learners ' Self- Regulation , Resilience And Autonomy*. <https://doi.org/10.1111/Ejed.70127>
- Zimmerman, B. (2000). *Attaining Self-Regulation: A Social Cognitive Perspective*. In M. Boekaerts, P. R. Pintrich, & M. Zeidner (Eds.) (Handbook O). Academic Press.